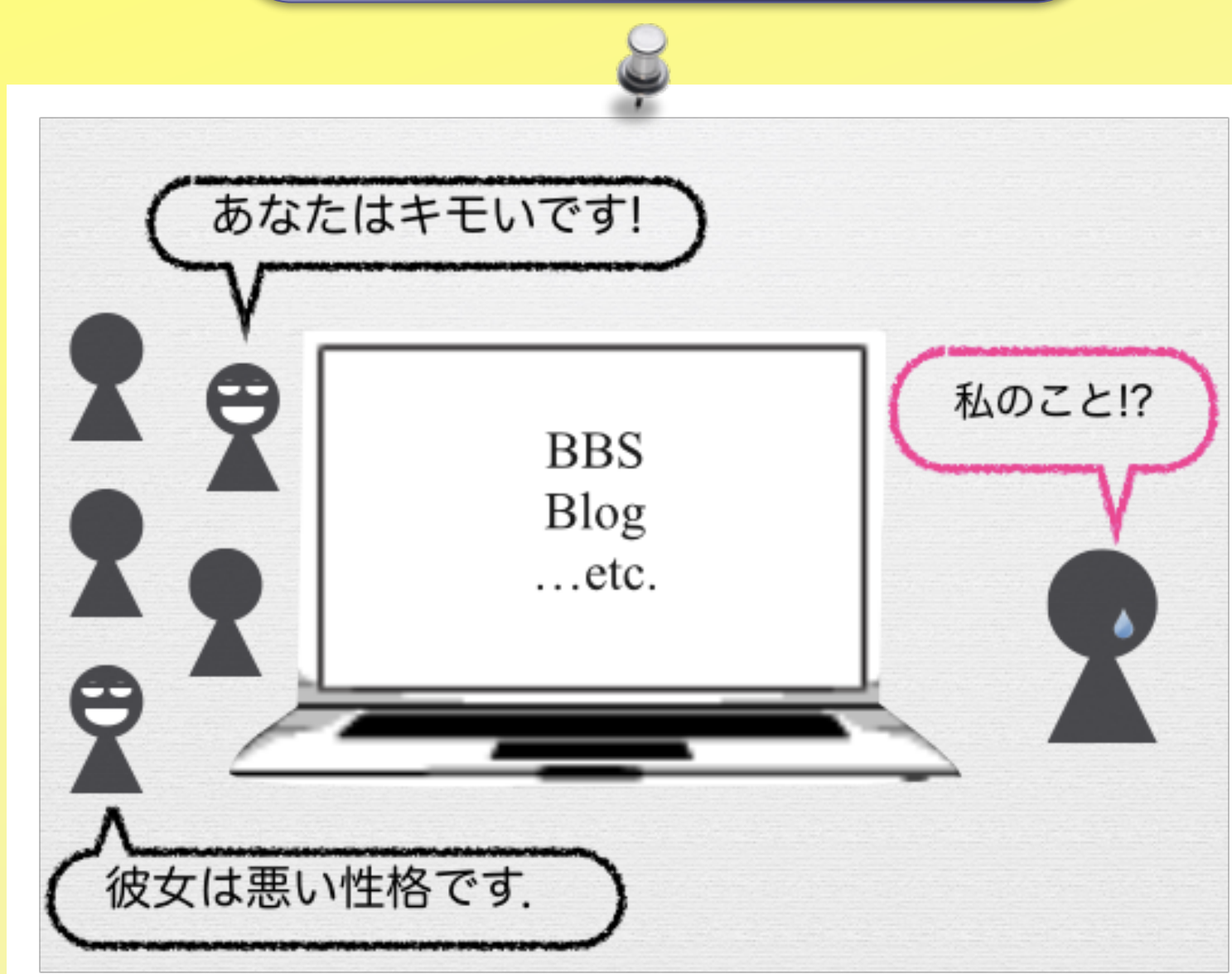


研究背景

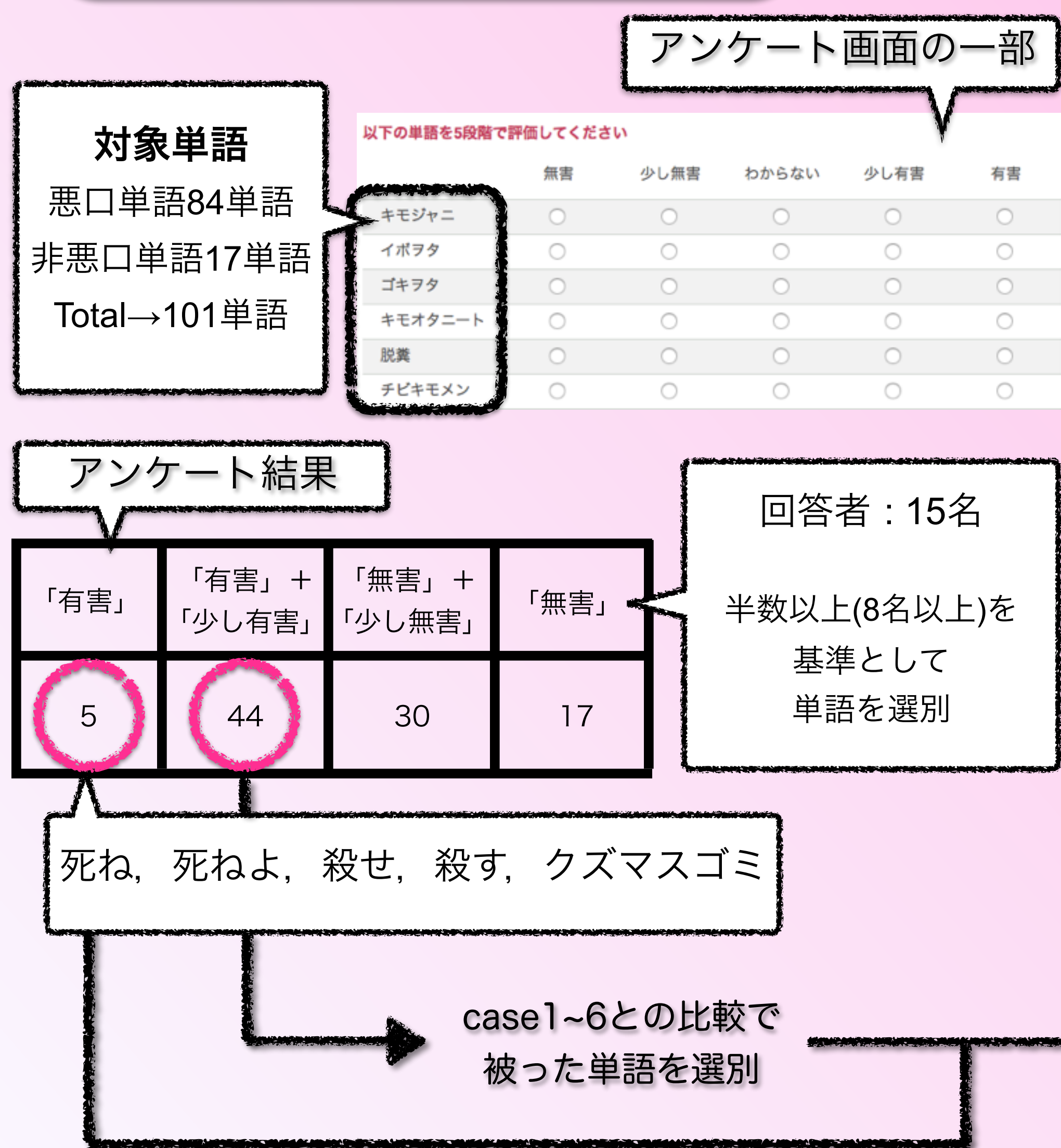


ネットいじめ
最近の社会問題

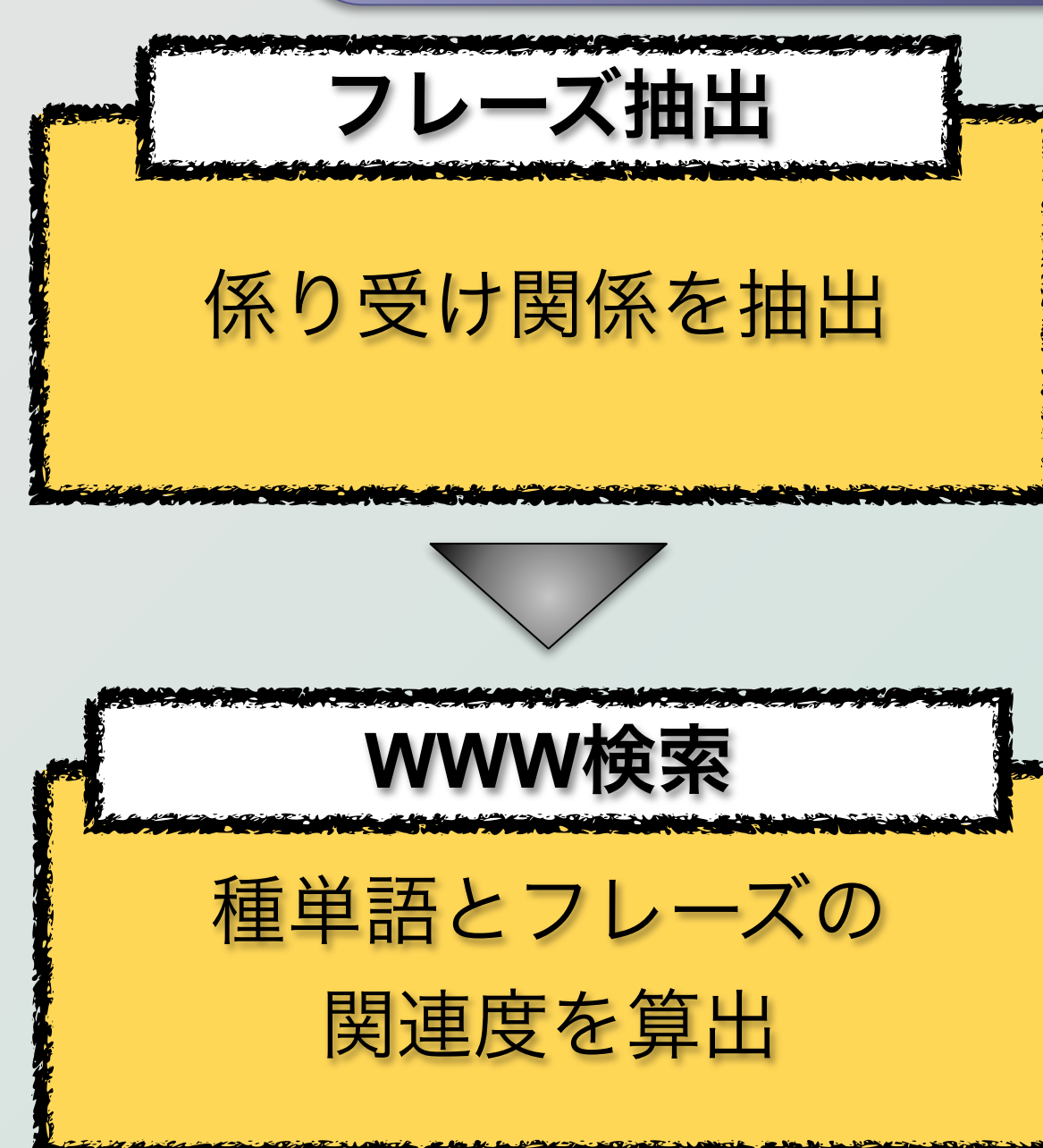
ネット
パトロール

- 学校関係者やPTAによるインターネットの監視
- 膨大な書き込み
→多くの労力と時間が必要
- 特定の人物しか閲覧できないSNS
→ネットパトロール自体が困難

人間解析による精緻化



カテゴリ別関連度最大化手法[1]

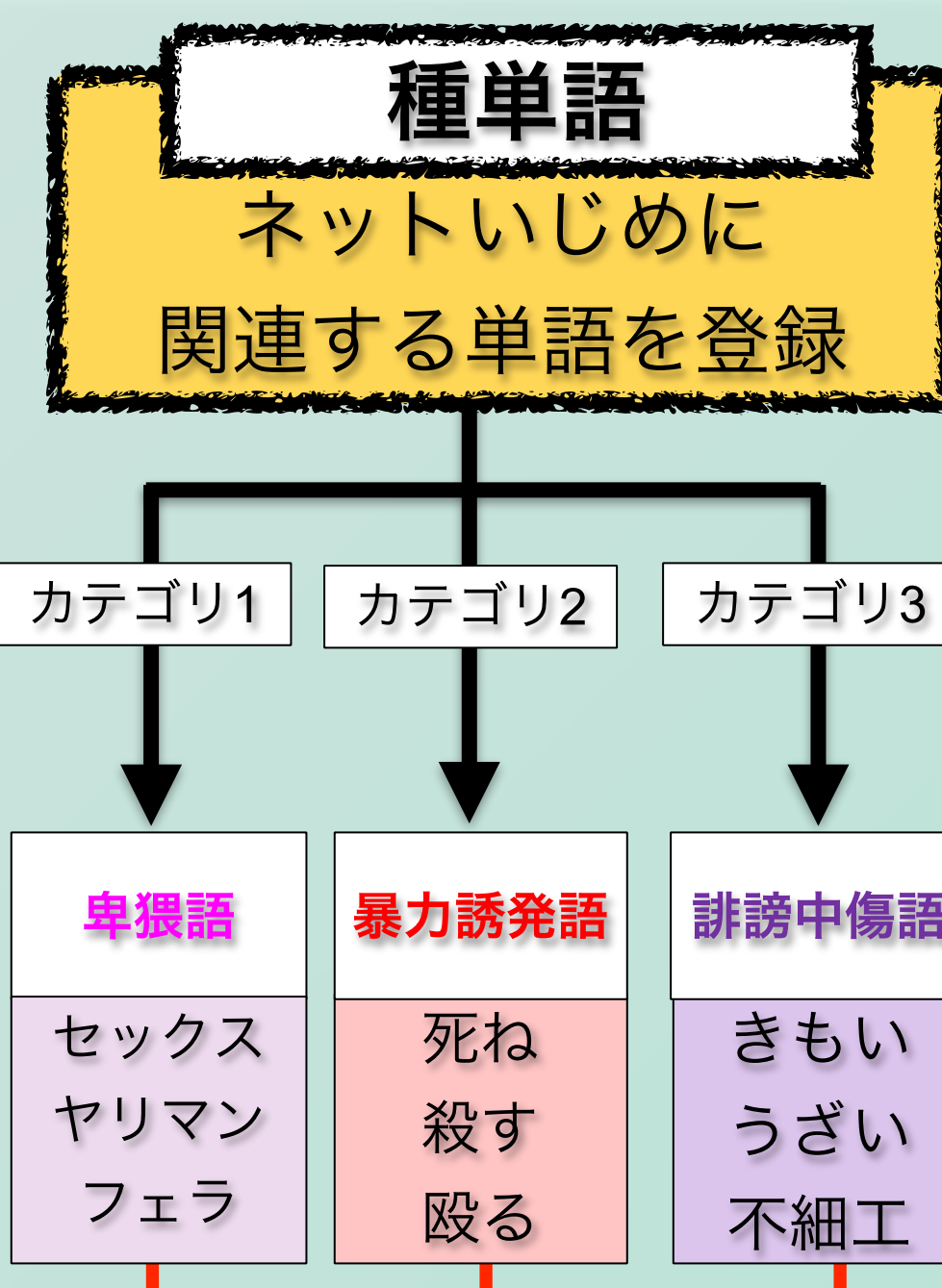


(名詞, 名詞) 例) 「性格が悪いかわいい女の子」
(名詞, 動詞) (女の子, かわいい), (性格, 悪い)
(名詞, 形容詞)

評価モデル(Turney'sの拡張SO-PMI [3])

$$score = \max (\max (PMI(p_i, w_j)))$$

フレーズ p_i と種単語 w_j の
AND検索により関連度を算出



本研究の位置付け

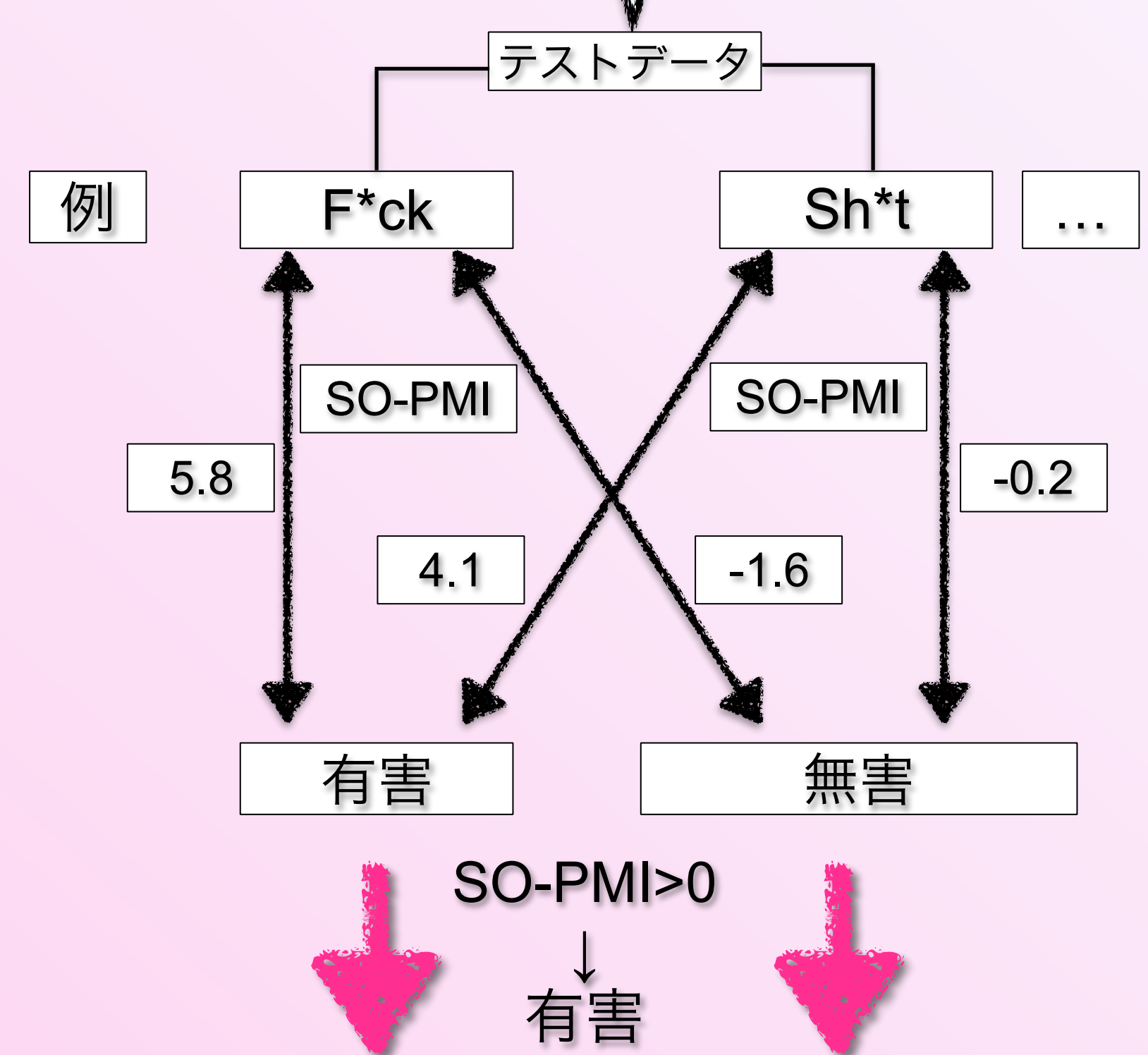
新田らの手法[1]の性能向上

種単語の規模や
組み合わせを考慮 → 種単語による
効果の検証

種単語の自動獲得[2]

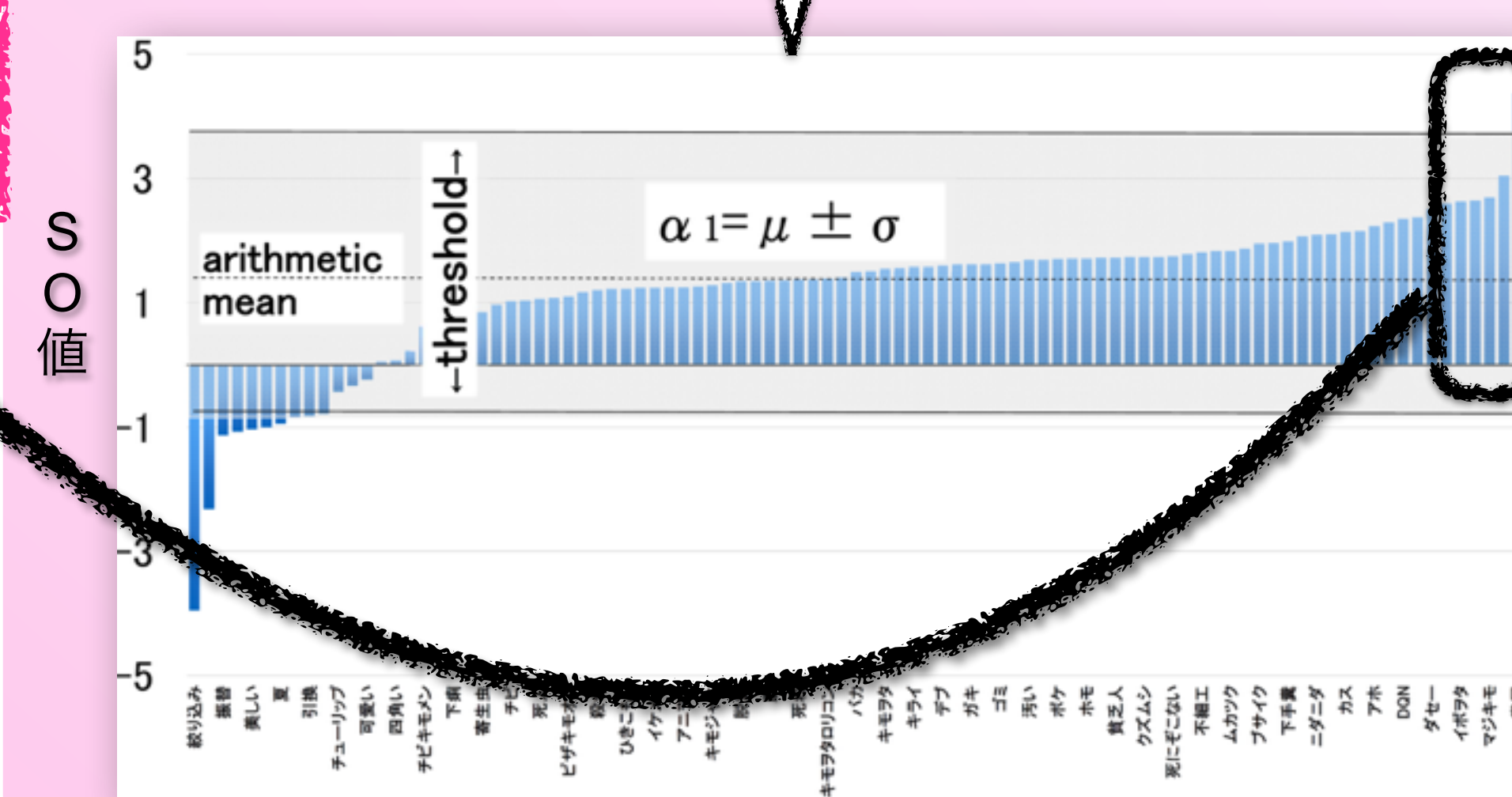
1次フィルタリング

新田らの種単語 × 石坂らの悪口単語 × 無害な単語



2次フィルタリング

石坂らの有害な種単語候補 × フィルタリングした種単語



テストデータ

比較実験

表1

	case1	case2	case3	case4	case5	case6
case5	0.676 vs 0.668 ***	0.676 vs 0.676 ***	0.676 vs 0.671 ***	0.676 vs 0.674 ***	-	0.676 vs 0.674 ***
case6	0.674 vs 0.668 ***	0.674 vs 0.676 ***	0.674 vs 0.671 ***	0.674 vs 0.674 ***	0.674 vs 0.676 ***	-

© Suzuha Hatakeyama 2016 * p<0.05, ** p<0.01, *** p<0.001

種単語の差による効果を検証

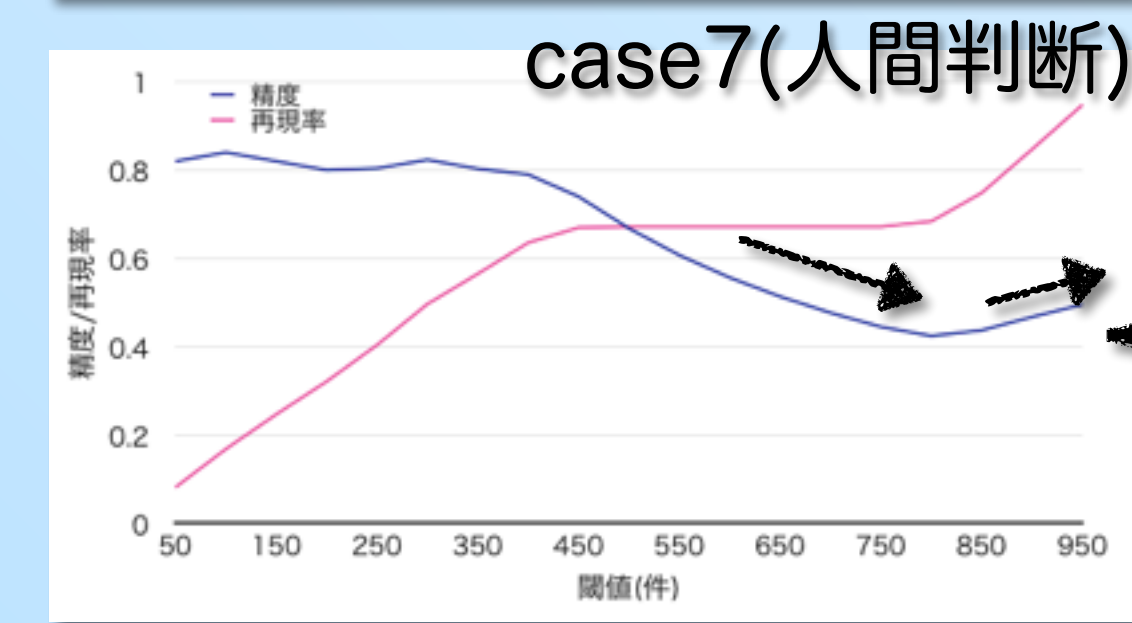
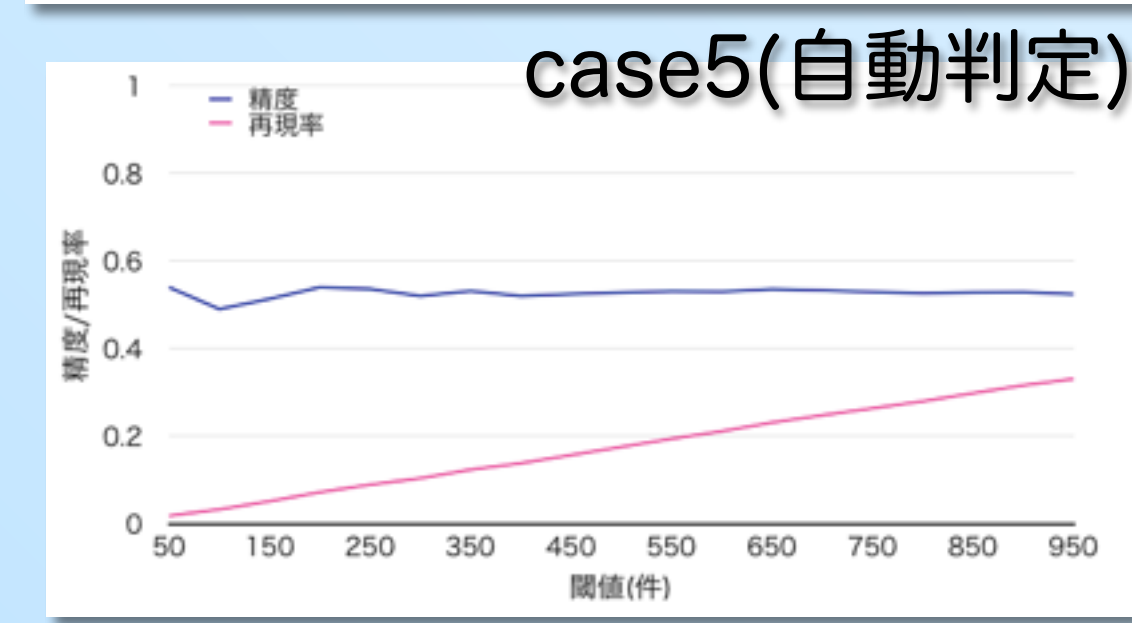
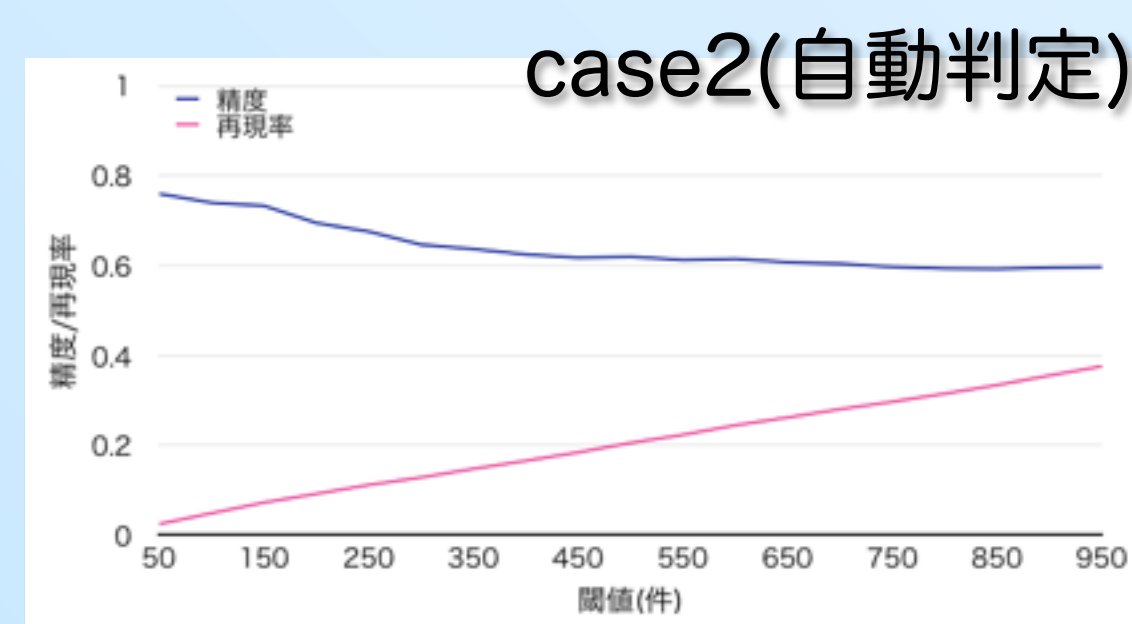
対象データ：テストデータ中の1975件
(phraseなしを除く)

マクネマー検定によりF値の有意差を確認
帰無仮説：F値に差がない(性能に差がない)

全てのcase(1~4)はcase6よりも性能が高い
case1と2はcase6よりも性能が高い
case5は最も性能が高い

テストデータ

有害：1508件
非有害：1490件
Total：2998件



文脈によって
有害にも
無害にも
変化する
単語が存在

自動検出

P 急激に減少

R 単調増加

人間判断

閾値150付近より急減少
800付近で上昇

閾値450付近まで単調増加,
450~850付近まで一定に
850付近から急上昇

まとめ

種単語の規模は性能に影響するかどうかの検証
→種単語の規模より質が重要

比較実験より

→人手で判断した単語は高い極性値を保持
→精度に影響を及ぼす

今後の課題

文脈により有害にも非有害にもなる文への対応

参考文献

- 新田大征, 榎井文人, プタシンスキ・ミハウ, 木村泰知, ジェブカ・ラファウ, 荒木健治: "カテゴリ別関連度最大化手法に基づく学校非公式サイトの有害書き込み検出", 第27回人工知能学会全国大会発表論文集, (2013.6).
- 石坂達也, 山本和英: "Web上の誹謗中傷を表す文の自動検出", 言語処理学会第17回年次大会発表論文集, pp.131-134(2011.3).
- Peter D. Turney, 2002. Thumbs Up or Thumbs Down? Semantic Orientation Applied to Unsupervised Classification of Reviews. In Proc. of ACL-2002, Philadelphia, pp.417-424, 2002.